

Testiranje statističke značajnosti/Testiranje statističkih hipoteza: rešenja zadataka sa vežbi

Rešenja uradio Andrej Gađanski, student psihologije, broj indeksa PS160059

Zadatak 1.

U fajlu **IQ.sav** nalaze se, između ostalog, (u varijabli **al4**) izvorni, tj. "sirovi" rezultati uzorka iz studentske populacije na testu za merenje inteligencije AL4 (autor testa: Konstantin Momirović).

- Izračunati IQ svakog ispitanika na osnovu sirovog rezultata ako se zna da je aritmetička sredina populacije za sirovi rezultat na ovom testu 25.71, a standardna devijacija 9.24. Varijabli dati ime **IQal4**;
- Napraviti 95% interval poverenja za aritmetičku sredinu na varijabli **al4**;
- Napraviti 99% interval poverenja za aritmetičku sredinu na varijabli **al4**;
- Statistički proveriti da li uzorak pripada populaciji u kojoj je prosečni IQ jednak 100 (koristiti već postojeću varijablu **IQ**);
- Statistički proveriti da li uzorak pripada populaciji u kojoj je prosečni IQ jednak 130 (koristiti već postojeću varijablu **IQ**).

Rešenje:

*Prvo je bitno napomenuti da ćemo u rešenjima zadataka postavljati nulte (i alternativne) hipoteze, a zatim analizirati da li imamo osnova za njihovo odbacivanje. Nivo značajnosti za odbacivanje nulte hipoteze, koji ćemo koristiti je 0.05. Podsetimo se da odabrani nivo značajnosti predstavlja zapravo verovatnoću da napravimo grešku tipa I (ova se verovatnoća uobičajeno označava oznakom α), tj. da odbacimo tacnu nultu hipotezu. Znači, ukoliko je verovatnoća p (koju dobijamo na osnovu rezultata sa uzorka i nulte distribucije uzorkovanja statistika za testiranje nulte hipoteze) veća od α (u našem slučaju 0.05) ne treba da odbacimo nultu hipotezu. Bitno je da se napomene da u slučaju da ne možemo da odbacimo nultu hipotezu to ne znači da je prihvatamo. Ako nam je npr. nulta hipoteza bila da je neka varijabla jednako izražena u 2 populacije i iz analize vidimo da ne treba odbaciti nultu hipotezu, ne smemo na osnovu tog jednog istraživanja zaključiti da ne postoje razlike u izraženosti te varijable u 2 populacije (to što nismo odbacili nultu hipotezu nije dovoljan argument za tvrdnju da nema razlike između tih populacija-**jer nepostojanje dokaza nije dokaz nepostojanja**). Jedino što možemo reći na osnovu datog istraživanja jeste da nije dobijena statistički značajna razlika.*

Pređimo na rešavanje 1. zadatka. Kako bismo izračunali IQ iz "sirovih" rezultata prvo treba da pronađemo standardizovane (z) skorove za svakog ispitanika, pri čemu ćemo za standardizaciju koristiti aritmetičku sredinu i standardnu devijaciju populacije. Pošto imamo μ i σ onda z-skorove lako možemo izračunati. Idemo u **Transform->Compute Variable**, u polje **Target variable:** možemo upisati **zal4** i u okvir **Numeric Expression:** unesemo **(al4-25.71)/9.24**. Ovime smo dobili varijablu koja sadrži z skorove. Tu varijablu ćemo transformisati u IQ za svakog ispitanika. Opet ćemo otići u **Compute Variable**. Novu varijablu ćemo u polju **Target variable:** nazvati **IQal4**. Pošto znamo da je aritmetička sredina z-skorova jednaka 0, a standardna devijacija jednaka 1 a pošto za IQ skalu važi da je $M = 100$ i $SD = 15$, IQ skorove ćemo dobiti formulom **100 + 15 * z**. Tako dobijamo novu varijablu u kojoj su ispisani IQ rezultati na osnovu testa al4.

Dakle, z skalu smo preveli u IQ skalu linearnom transformacijom u kojoj je aditivna konstanta 100 a multiplikativna konstanta 15. To sledi na osnovu sledećih jednakosti:

$M_t = a + b * M$ (aritmetička sredina se pri linearnim transformacijama menja za aditivnu konstantu i za multiplikativnu konstantu puta)

$S_t = b * S$ (standardna devijacija se pri linearnim transformacijama menja za multiplikativnu konstantu puta)

pri čemu su M_t i S_t aritmetička sredina i standardna devijacija rezultata posle linearne transformacije, a M i S aritmetička sredina i standardna devijacija rezultata pre linearne transformacije.

Prema tome, znamo da M_t treba da bude 100 a S_t treba da bude 15. Budući da je $S_t = 15$, a $S = 1$, multiplikativnu konstantu odredićemo zamenom S_t i S u drugoj jednakosti: $15 = b * 1$, tj. $b = 15$. Pošto znamo da je $M = 0$, $S = 1$ a $b = 15$, aditivnu konstantu ćemo odrediti zamenom M_t , M i b u prvoj jednakosti: $100 = a + 15 * 0$, tj. $a = 100$.

Napomena: Transformaciju u IQ vrednosti mogli smo napraviti linearnom transformacijom i direktno iz sirovih rezultata ali bi tada konstante a i b bile znatno složenije za računanje.

Da bismo izvorne rezultate direktno preveli u skalu čija je aritmetička sredina 100, a standardna devijacija 15 koristimo se linearnom transformacijom oblika

$$y = a + b * x$$

pri čemu a i b predstavljaju aditivnu i multiplikativnu konstanta čije je vrednosti potrebno ustanoviti.

Kao i u prethodnom slučaju, postavljamo sistem od dve jednačine sa dve nepoznate:

$$1. S_t = b * S$$

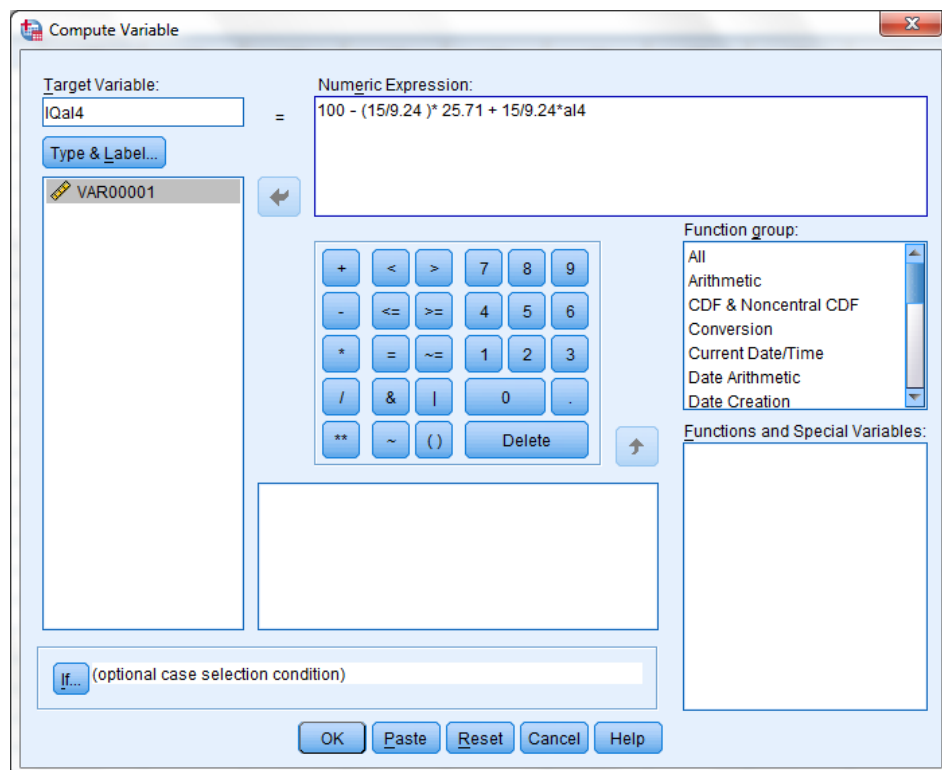
$$2. M_t = a + b * M$$

i njihovim rešavanjem dobijamo:

$$1. b = S_t / S, \text{ tj. } b = 15 / 9.24$$

$$2. a = M_t - b * M \text{ ili } a = M_t - (S_t / S) * M, \text{ tj. } a = 100 - (15 / 9.24) * 25.71$$

što znači da bi prozor komande Compute izgledao ovako:



Napomena: Vrednosti na varijabli **IQal4** nisu iste kao vrednosti na varijabli **IQ** koja već postoji u fajlu. To je zato što se IQ na ovoj bateriji testova (koju pored al4 čine još dva testa) zapravo računa na osnovu rezultata na tri testa. (Postupak računanja IQ-a na bateriji KOG3 dat je u priručniku baterije i nećemo ga ovde prikazivati).

Sada treba napraviti 95% interval poverenja za aritmetičku sredinu na varijabli **al4**. To možemo uraditi tako što ćemo otići u **Analyze->Descriptive Statistics->Explore**, u okvir **Dependent list:** ćemo prebaciti varijablu **al4**. Pošto nam grafici ovde nisu neophodni, u

okviru "Display" biramo "Statistics". Zatim, klikom na dugme **Statistics** videćemo da je već odabran interval 95%. Zato kliknemo na **Continue->OK**.

U ispisu dobijamo sledeću tabelu:

Descriptives			Statistic	Std. Error
al4	Mean		38.41	.180
	95% Confidence Interval for Mean	Lower Bound	38.05	
		Upper Bound	38.77	
	5% Trimmed Mean		38.64	
	Median		39.00	
	Variance		4.838	
	Std. Deviation		2.200	
	Minimum		23	
	Maximum		40	
	Range		17	
	Interquartile Range		3	
	Skewness		-2.854	.199
	Kurtosis		15.219	.395

Na osnovu dobijenog intervala možemo tvrditi sa sigurnošću od 95% da je parametar obuhvaćen tim intervalom. Odnosno, možemo sa sigurnošću od 95% tvrditi da je aritmetička sredina rezultata na testu al4 u populaciji između 38.05 i 38.77. Dakle, postoji rizik od 5% da 95% intervalom poverenja nismo obuhvatili parametar.

Postupak za dobijanje 99% intervala pouzdanosti je analogan upravo prikazanom. Jedino se u **Explore: Statistics** mora ukucati "99" pošto je po defaultu tu upisan interval od 95%. Vidimo da je 99% interval širi pošto je verovatnoća da slučajna varijabla uzme vrednost izvan njega jednaka 0.01, tj. 1%.

Za sledeći deo zadatka postaviti ćemo sledeću nultu hipotezu:

$$H_0: \mu=100$$

Da bismo testirali ovu hipotezu možemo upotrebiti proceduru **Analyze->Compare Means->One-Sample T Test**. U ovom prozoru ćemo prebaciti varijablu **IQ** u okvir **Test variable(s):**, a u polju **Test value:** ćemo ukucati **100**. Time ćemo u ispisu dobiti sledeće tabele:

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
Ukupni IQ	149	118.05	14.125	1.157

One-Sample Test

	Test Value = 100					
	t	df	Sig. (2-tailed)	Mean Difference	95% Confidence Interval of the Difference	
					Lower	Upper
Ukupni IQ	15.602	148	.000	18.054	15.77	20.34

Statistik t predstavlja odstupanje aritmetičke sredine uzorka (u ispisu **Mean**) od pretpostavljene aritmetičke sredine populacije (odstupanje je u ispisu **Mean Difference**) podeljeno sa standardnom greškom za aritmetičku sredinu (u ispisu **Std. Error Mean**). U koloni **Sig.(2-tailed)** data je verovatnoća da t statistik uzme vrednost 15.602 ili veću, bez obzira na predznak (dakle ili veću od 15.602 ili manju od -15.602 ili jednaku ovim vrednostima) **ako je nulta hipoteza tačna**. Vidimo da je ta verovatnoća manja od 0.05 što nam daje osnova da odbacimo nultu hipotezu. Takođe, na drugi način ovaj zaključak možemo iskazati tako što ćemo reći da je t statistik statistički značajan ili da je odstupanje aritmetičke sredine uzorka od pretpostavljenog parametra statistički značajno. Praktično, ovo nam govori da uzorak ne pripada populaciji u kojoj je prosečan IQ jednak 100.

Napomena: Često će u koloni **Sig.(2-tailed)** pisati .000. To ne znači da je ova verovatnoća zaista 0 nego da se tek na nekom daljem decimalnom mestu koje se ne vidi nalazi neka cifra koja nije 0.

Postupak za sledeći deo zadatka je potpuno isti samo sa drugačijom vrednošću u polju **Test value:** (130), te ga neću analizirati. Na kraju će se dobiti nešto manji t statistik ali ćemo i ovde imati osnova za odbacivanje nulte hipoteze.

Kako bismo pojasnili šta zapravo to **Sig.(2-tailed)** [sig je skraćeno od significance, dok 2 tailed pokazuje da je reč o dvosmernom testiranju] citiraćemo profesora Tenjovića:

Dvosmerno testiranje: p [Sig. (2-tailed)] je verovatnoća da, AKO JE NULTA HIPOTEZA TAČNA, statistik t koji služi za testiranje nulte hipoteze uzme (BEZ OBZIRA NA PREDZNAK) vrednost jednaku ili veću od one vrednosti koju smo dobili na slučajnom uzorku. Tu verovatnoću računamo iz nulte distribucije uzorkovanja t -statistika, tj. distribucije uzorkovanja koju ima t statistik AKO JE NULTA HIPOTEZA TAČNA.

$$[\text{Sig. (2-tailed)}] = p = P(|t| \geq t_{\text{dobijeno}} \mid H_0 \text{ tačna})$$

Zadatak 2.

U fajlu **astma.sav** postoje podaci o neuroticizmu (varijabla **EPQN**) za astmatičare i neastmatičare (varijabla **GRUPA** pokazuje kojoj od ovih kategorija ispitanik pripada).

- Testirajte pretpostavku o normalnosti distribucije neuroticizma u populaciji astmatičara.
- Da li se astmatičari i neastmatičari u populaciji razlikuju prema prosečnom neuroticizmu?
- Napraviti 95% interval poverenja za razliku između aritmetičkih sredina astmatičara i neastmatičar? Šta na osnovu intervala poverenja možemo zaključiti?

Rešenje:

Ovde ćemo postaviti nultu hipotezu da se neuroticizam normalno distribuira u populaciji astmatičara. Dakle:

$H_0: X \sim N(\mu = M ; \sigma = S)$ [Varijabla X ima normalnu raspodelu u populaciji sa parametrima μ i σ jednakim aritmetičkoj sredini i standardnoj devijaciji dobijenim na uzorku]

$H_1: X \text{ nije } \sim N$

Postoji više načina na koje možemo testirati ovu hipotezu, ali ćemo to najlakše uraditi sa dva testa normalnosti: Lillieforsovom modifikacijom Kolmogorov-Smirnovljevog testa i Šapiro-Wilkovim testom (ovaj drugi se može koristiti i na malim uzorcima – samo moraju biti veći od 30).

*Ali prvo moramo podeliti naše ispitanike na grupe astmatičara i neastmatičara kako bismo obradili podatke samo za astmatičare. Za to idemo u meni **Data** pa na **Select cases**. Tu biramo **If condition is satisfied**: i kliknemo na dugme **IF** kako bismo postavili uslov. Ovim dobijamo prozor sličan onom u komandi Compute. U polje **Numeric expression**: prebacujemo varijablu **Grupa** i napišemo da je jednaka 1 (u variable view smo mogli videti da je 1 oznaka za astmatičare, a 2 za neastmatičare). Onda idemo na **Continue** pa na **OK** i sada, sve dok ne isključimo ovaj uslov, SPSS će sve analize i grafike raditi samo za odabranu grupu ispitanika.*

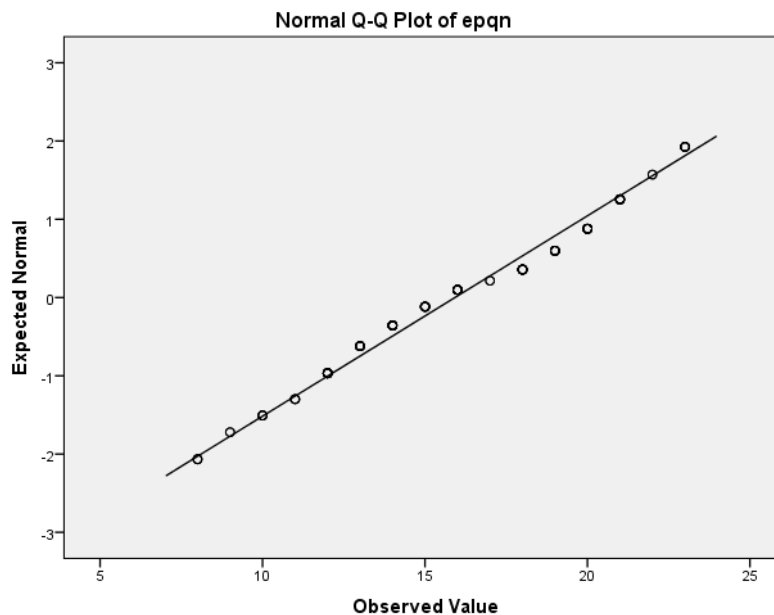
*Sada idemo na **Analyze->Descriptive Statistics->Explore**. Da bismo uradili statističke testove, u donjem delu novodobijenog prozora (u okviru "Display") selektujemo **Plots** i onda kliknemo na dugme **Plots** u desnom delu prozora. Tu ćemo štriklirati **Normality plots***

with tests i onda idemo na **Continue** i **OK**. U ispisu ćemo dobiti sledeću tabelu:

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
epqn	.109	128	.001	.965	128	.002

a. Lilliefors Significance Correction

Pošto zaključujemo na nivou značajnosti od 0.05, ukoliko je **Sig.** < 0.05 imamo osnova za odbacivanje nulte hipoteze. Kao što vidimo, na oba testa se dobija verovatnoća *p* manja od 0.05 te možemo reći da treba da odbacimo nultu hipotezu o normalnoj raspodeli neuroticizma među astmatičarima. Ovo sve možemo videti i na **QQ Plot (Quantile-Quantile plot)** dijagramu koji takođe dobijamo u ispisu procedure koju smo izveli. Na dijagramu su prikazani kvantili empirijske raspodele (X-osa) naspram kvantila koji se očekuju za normalnu raspodelu (Y-osa).



Da je distribucija normalna, svi ovi markeri bi ležali na liniji, ali vidimo da postoje odstupanja, koja jesu mala ali su statistički značajna.

Pre nego što započnemo drugi deo zadatka moramo isključiti prethodno postavljeni uslov u **Select Cases** kako SPSS ne bi sve radio samo za astmatičare. Sada ćemo tamo uključiti opciju **All Cases**.

Za drugi deo zadatka ćemo postaviti sledeću nultu i alternativnu hipotezu:

$H_0: \mu_1 = \mu_2$ (ne postoji razlika po prosečnom neroticizmu između astmatičara i neastmatičara u populaciji)

$H_1: \mu_1 \neq \mu_2$

Ali prvo da vidimo možemo li koristiti t-test ako imamo razloga da sumnjamo u pretpostavku o normalnoj distribuciji, kao što je to u našem slučaju. Naime, prethodno smo barem za jednu od ovih dveju populacija odbacili ovu pretpostavku. Tu se možemo pozvati na određeni stepen robustnosti t-testa u odnosu na odstupanja od normalne distribucije: drugim rečima, t-test će na ovoliko velikim uzorcima funkcionisati dobro i ako varijabla nije normalno distribuirana u svakoj populaciji.

Nultu hipotezu ćemo proveriti t-testovima. Prvo moramo definisati kakvi su naši uzorci, zavisni ili nezavisni. Oni mogu biti nezavisni ako odabir člana u jedan uzorak ni na koji način ne utiče na verovatnoću odabira člana u drugi uzorak. To je, na primer, slučaj ako testiramo razlike između muškaraca i žena i za svakog slučajno odabranog muškarca uzimamo slučajno odabranu ženu. Drugi slučaj je kada odabir jednog člana u uzorak utiče na verovatnoću odabira člana u drugi uzorak. Tada je reč o zavisnim uzorcima. To bi bio slučaj kada bi za svakog nasumično odabranog muškarca automatski birali njegovu ženu u uzorak žena. Zavisni uzorci u psihološkim istraživanjima postoje i kada se isti ispitanici testiraju na više varijabli ili se isti ispitanici ispituju u više situacija.

U ovom zadatku su u pitanju dva nezavisna uzorka. Zato, kako bismo proverili naše hipoteze, možemo izvršiti sledeću proceduru: **Analyze->Compare Means->Independent Samples T-test**. Pošto želimo da istražimo razlike u prosečnom neuroticizmu, kao **Test variable**: unosimo **EPQN** a kao **Grouping variable** ubacujemo **GRUPA**. Potom moramo definisati kako su u podacima označene grupe, pa kliknemo na **Define Groups**: i tu ukucamo oznake za naše grupe. U ovom slučaju to će biti 1 i 2 jer u variable view možemo videti da su astmatičari označeni jedinicom, a neastmatičari dvojkom. To je sve što treba da uradimo i sada u ispisu dobijamo sledeće tabele:

Group Statistics

grupa	N	Mean	Std. Deviation	Std. Error Mean
epqn astmaticari	128	15.92	3.906	.345
neastmaticari	128	10.77	3.842	.340

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
epqn	Equal variances assumed	.251	.617	10.632	254	.000	5.148	.484	4.195	6.102
	Equal variances not assumed			10.632	253.931	.000	5.148	.484	4.195	6.102

U donjoj tabeli vidimo da imamo dva reda, i u oba reda su rezultati t-testova, ali oni ne moraju uvek biti isti. U testu iz prvog reda imamo takozvanu pretpostavku o homoscedastičnosti, tj. pretpostavku o

tome da su varijanse populacija jednake, dok u drugom redu to nije pretpostavljeno. Dakle, t-test u prvom redu valjan je samo ako postoji homoscedastičnost. Da bismo znali koju varijantu t-testa ćemo koristiti pogledaćemo verovatnoću u koloni **Sig.** u Leveneovom testu i ukoliko je ova verovatnoća veća od 0.05 onda možemo koristiti test u prvom redu, a ako je jednaka 0.05 ili manja onda moramo koristiti test u drugom redu. U našem slučaju ova verovatnoća iznosi 0.617 pa možemo koristiti test koji je prikazan u prvom redu.

E sada, šta možemo zaključiti o našim hipotezama? Vidimo da nam je t statistik 10.632 sa 254 stepena slobode. Da bismo izveli zaključak o značajnosti ovog statistika moramo pogledati verovatnoću u koloni **Sig.(2-tailed)** i ukoliko je ona manja od 0.05 možemo reći da je t-statistik statistički značajan, tj. da treba da odbacimo nultu hipotezu. Ovo znači u našem slučaju da zaključujemo da postoji razlika u prosečnom neuroticizmu između astmatičara i neastmatičara u populaciji.

Za poslednji deo ovog zadatka, pogledaćemo u 95% interval poverenja koji smo dobili automatski u ispisu (kolona **95% Confidence interval of the Difference**) tokom prethodne procedure. Inače, ovaj interval definišemo klikom na dugme **Options** prilikom pravljenja T testa sa nezavisnim uzorcima (po defaultu on je 95%). Vidimo da su granice ovog intervala 4.195 i 6.102. Na osnovu toga možemo reći sa sigurnošću od 95% da je razlika između aritmetičkih sredina neuroticizma kod astmatičara i neastmatičara u populaciji obuhvaćena ovim intervalom. Takođe primećujemo da ovaj interval ne obuhvata 0. Ukoliko bi obuhvatao 0 razlika između aritmetičkih sredina ne bi bila statistički značajna. Ovako vidimo da jeste statistički značajna na nivou značajnosti od 0.05.

Zadatak 3.

Fajl sa podacima: **W.sav**. U fajlu su nalaze mere studenata na visini i težini (varijabla **WEIGHT**), kao i informacija o polu (varijabla **SEX**).

- Utvrditi da li u populaciji postoji razlika u prosečnoj težini muških i ženskih studenata?

Rešenje:

Postupak za rešavanje ovog zadatka je potpuno isti kao za drugi deo prethodnog zadatka.

Za početak ćemo postaviti sledeće hipoteze:

$H_0: \mu_1 = \mu_2$ (ne postoji razlika u populaciji između muških i ženskih studenata po prosečnoj težini)

$H_1: \mu_1 \neq \mu_2$

U ovom zadatku su u pitanju dva nezavisna uzorka. Zato, kako bismo testirali nultu hipotezu opet ćemo izvesti istu proceduru: **Analyze->Compare Means->Independent Samples T-test**. Pošto želimo da istražimo razlike u prosečnoj težini, u okvir **Test variable** unosimo **WEIGHT** a kao **Grouping variable** biramo **SEX**. Potom kliknemo na **Define Groups**, unesemo 1 i 2 jer u **Variable view** možemo videti da su muškarci označeni jedinicom, a žene dvojkom. Kliknemo na **OK** i dobijamo prozor za ispis sa sledećim tabelama:

Group Statistics					
sex or gender [begins the Confidential Personal Information measure]		N	Mean	Std. Deviation	Std. Error Mean
wt in kilograms	Male	100	76.87	8.754	.875
	Female	100	58.44	7.189	.719

Independent Samples Test										
		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
wt in kilograms	Equal variances assumed	3.998	.047	16.269	198	.000	18.430	1.133	16.196	20.664
	Equal variances not assumed			16.269	190.784	.000	18.430	1.133	16.196	20.664

E sada, šta možemo zaključiti o našim hipotezama? Budući da imamo razloga da sumnjamo u pretpostavku homoscedastičnosti (verovatnoća u koloni **Sig.** Levenovog testa iznosi 0.047, tj. manja je od 0.05) gledaćemo ishod t-testa u drugom redu. Vidimo da nam je t statistik 16.269 sa 190.784 stepeni slobode. Da bismo izveli zaključak o značajnosti ovog statistika moramo pogledati verovatnoću u koloni **Sig.(2-tailed)** i pošto je ona manja od 0.05 možemo reći da je t-statistik statistički značajan, tj. da treba da odbacimo nultu hipotezu. Ovo znači da zaključujemo da postoji razlika u prosečnoj težini muških i ženskih studenata u populaciji.

Zadatak 4

Podaci su u fajlu **Nikomal.sav**. U nameri da proveriti delotvornost nove žvakaće gume, psiholog u firmi koja proizvodi žvakaću gumu Nikomal je u slučajnom uzorku pušača snimio broj cigareta koje popuše tokom nedelju dana (varijabla **nkm_pre**). Zatim je svaki ispitanik iz ovog uzorka dobio pakovanja žvakaće gume i uputstvom kada i kako da je koriste. Isto tako, ispitanici su tokom nedelju dana u kojoj su koristili Nikomal snimili broj popušanih cigareta (varijabla **nkm_guma**).

- Do kakvog zaključka je psiholog došao: da li podaci sugerišu da je korišćenje žvakaće gume Nikomal imalo uticaja na broj popušanih cigareta?

Rešenje:

Na prvi pogled vidimo da se broj popušenih cigareta nakon korišćenja žvaki smanjio u uzorku u odnosu na period pre upotrebe „Nikomala”. Postoji li efekat žvakaće gume u populaciji možemo proveriti tako što ćemo opet postaviti par početnih hipoteza za analizu:

$H_0: \mu_d = 0$ (μ_d je aritmetička sredina razlika parova rezultata u populaciji)

$H_1: \mu_d \neq 0$

Ovo praktično znači da pretpostavljamo da je broj prosečno popušenih cigareta bez nikomala jednak broju popušenih cigareta uz korišćenje nikomala. Ovde lako vidimo da su u pitanju zavisni uzorci, pošto na istim ispitanicima procenjujemo dva različita parametra, odnosno meru popušenih cigareta sa i bez tretmana.

Za testiranje nulte hipoteze upotrebićemo proceduru **Analyze->Compare Means->Paired Samples T-Test**. Ovde ćemo naše 2 varijable prebaciti na desnu stranu u jedan par i zatim kliknuti na **OK**. U ispisu ćemo dobiti tri tabele:

Paired Samples Statistics					
		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	Broj popušenih cigareta u nedelji bez Nikomala	148.71	14	43.059	11.508
	Broj popušenih cigareta u nedelji sa Nikomalom	137.79	14	48.286	12.905

Paired Samples Correlations			
	N	Correlation	Sig.
Pair 1 Broj popušenih cigareta u nedelji bez Nikomala & Broj popušenih cigareta u nedelji sa Nikomalom	14	.952	.000

Paired Samples Test									
		Paired Differences				t	df	Sig. (2-tailed)	
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower				Upper
Pair 1	Broj popušenih cigareta u nedelji bez Nikomala - Broj popušenih cigareta u nedelji sa Nikomalom	10.929	15.117	4.040	2.200	19.657	2.705	13	.018

Iz prve tabele vidimo da u uzorku postoji razlika između aritmetičkih sredina, te da je broj popušenih cigareta manji u proseku nakon upotrebe nikomala. U drugoj tabeli vidimo koeficijent korelacije broja popušenih cigareta u dve situacije. Bitno je da ovaj koeficijent ne iznosi 0 ili da nije veoma blizu nuli, jer onda ne bi baš imalo smisla raditi t-test. U trećoj tabeli vidimo razliku između aritmetičkih sredina i ostale podatke. Statistika t je 2.705 i on je zapravo aritmetička sredina razlika između broja popušenih cigareta u dve situacije podeljena standardnom greškom za tu aritmetičku sredinu. Vidimo da je broj stepeni slobode 13. To je zato što imamo 14 parova podataka pa je $df = n - 1$. Vidimo da je verovatnoća u koloni **Sig.(2-tailed)** manja od

0.05 što nam govori da treba da odbacimo nultu hipotezu, tj. da je t statistik statistički značajan.

Iz ovoga bi psiholog mogao da zaključi da u proseku, nikomal ima kakav-takav uticaj na smanjenje broja popušenih cigareta.

Zadatak 5.

Otvorite bilo koji fajl sa podacima (konkretni podaci za ovaj zadatak nisu bitni).

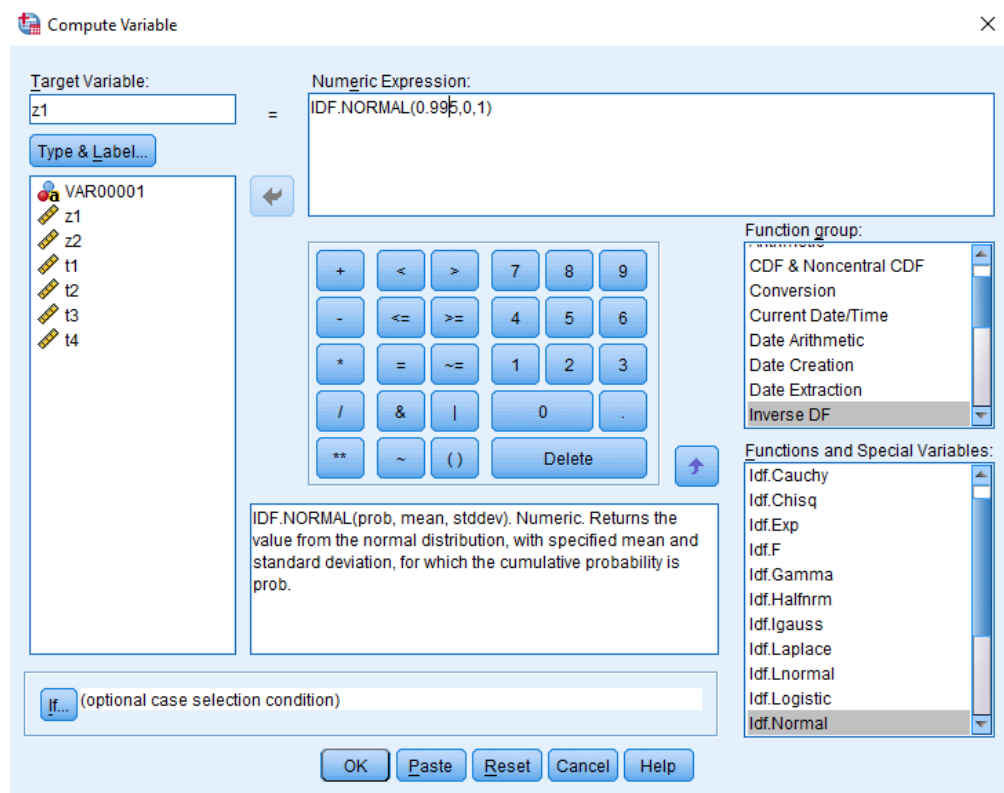
- Ako slučajna varijabla z ima standardizovanu normalnu raspodelu izračunajte vrednost **z1** za koju važi sledeće: verovatnoća da slučajna varijabla z uzme vrednost jednaku z1 ili veću od z1 jednaka je 0.005;
- Ako slučajna varijabla z ima standardizovanu normalnu raspodelu izračunajte vrednost **z2** za koju važi sledeće: verovatnoća da slučajna varijabla z uzme vrednost jednaku z2 ili manju od z2 jednaka je 0.005;
- Ako slučajna varijabla t (t statistik) ima Studentovu T raspodelu sa 18 stepeni slobode izračunajte vrednost **t1** za koju važi sledeće: verovatnoća da slučajna varijabla t uzme vrednost jednaku t1 ili veću od t1 jednaka je 0.005;
- Ako slučajna varijabla t (t statistik) ima Studentovu T raspodelu sa 18 stepeni slobode izračunajte vrednost **t2** za koju važi sledeće: verovatnoća da slučajna varijabla t uzme vrednost jednaku t2 ili manju od t2 jednaka je 0.005;
- Ako slučajna varijabla t (t statistik) ima Studentovu T raspodelu sa 998 stepeni slobode izračunajte vrednost **t3** za koju važi sledeće: verovatnoća da slučajna varijabla t uzme vrednost jednaku t3 ili veću od t3 jednaka je 0.005;
- Ako slučajna varijabla t (t statistik) ima Studentovu T raspodelu sa 998 stepeni slobode izračunajte vrednost **t4** za koju važi sledeće: verovatnoća da slučajna varijabla t uzme vrednost jednaku t4 ili manju od t4 jednaka je 0.005;
- Uporedite vrednost z1 sa vrednostima t1 i t3, a vrednost z2 sa vrednostima t2 i t4. Šta na osnovu toga možete da zaključite? Postoji li statistička teorema na osnovu koje se ovi rezultati mogu predvideti?

Rešenje:

*U ovom zadatku možemo otvoriti bilo koji fajl sa podacima ili u prazan prozor za podatke upisati bilo kakav podatak, na primer, cifru 1, kako bismo mogli da koristimo komandu Compute. Da bismo uradili ovaj zadatak korišćemo komandu **Transform->Compute** i funkcije **IDF.NORMAL** i **IDF.T**. **IDF.NORMAL** ćemo koristiti kada imamo normalnu raspodelu. Za ovu funkciju moramo definisati 3 vrednosti: verovatnoću, aritmetičku sredinu i standardnu devijaciju. Druga funkcija je **IDF.T** i nju ćemo koristiti kada imamo Studentovu*

raspodelu. Za ovu funkciju moramo definisati dve vrednosti: verovatnoću i broj stepeni slobode.

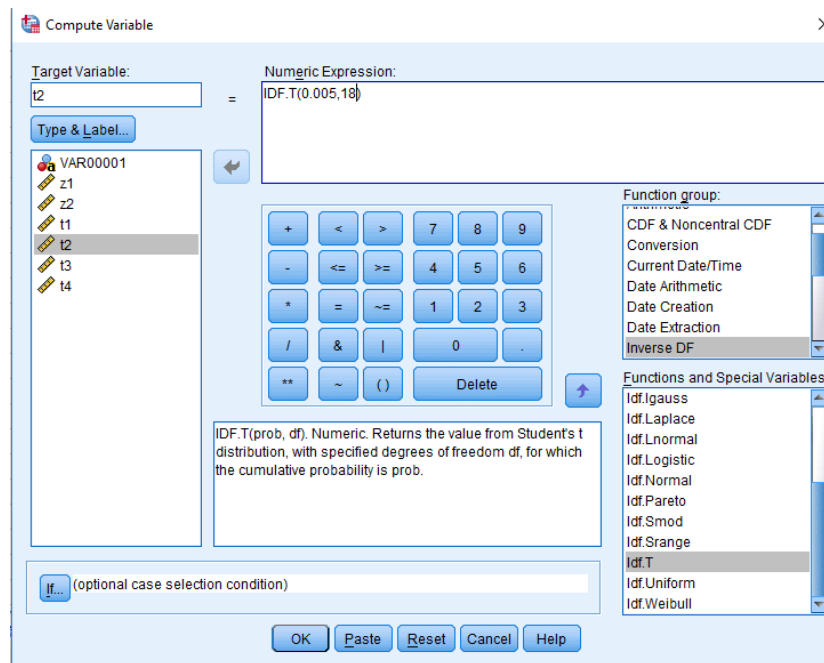
Dakle, kako bismo izračunali vrednost **z1** (tako ćemo i nazvati Target variable), nakon što odaberemo **Compute variable** među funkcijama naći ćemo **IDF.NORMAL** i uneti tražene vrednosti. Prva zahtevana vrednost je kumulativna verovatnoća do vrednosti z1. Pošto u zadatku piše da je verovatnoća da varijabla uzme vrednost istu ili veću od z1 jednaka 0.005 onda je verovatnoća da varijabla uzme neku vrednost manju od z1 jednaka **1 - 0.005**, tj. **0.995**. Tu vrednost ćemo i uneti u našu funkciju. Sada su nam potrebne i aritmetička sredina i standardna devijacija. Pošto iz zadatka vidimo da je u pitanju standardizovana normalna raspodela, znamo da je **M = 0** i **SD = 1**. Pa ćemo i te vrednosti uneti na tražena mesta.



Tako ćemo dobiti novu "varijablu" z i u njoj ćemo videti da je rešenje **2.58**.

Postupak za računanje **z2** je isti, jedino što ćemo na mesto kumulativne verovatnoće do vrednosti z2 uneti **0.005** (pošto je sad dato da je verovatnoća da varijabla uzme vrednost manju od z2 jednaka 0.005). Kao rešenje ćemo dobiti **-2.58**. Dakle, verovatnoća da slučajna varijabla koja ima standardizovanu normalnu raspodelu uzme neku vrednost u intervalu od -2.58 do +2.58 iznosi 0.99, ili 99%. Vrednost -2.58 predstavlja kvantil 0.005, a vrednost +2.58 predstavlja kvantil 0.995 na standardizovanoj normalnoj raspodeli.

Postupak za računanje kvantila za Studentovu t raspodelu, tj. traženih t vrednosti je gotovo isti, te nećemo ulaziti u detalje. U ovom slučaju koristimo funkciju **IDF.T** i to je jedina razlika u postupku. Ova funkcija takođe kao prvu vrednost zahteva da se unese kumulativna verovatnoća do vrednosti tražene u zadatku (za t_1 i t_3 ona će biti 0.995 a za t_2 i t_4 0.005 što se vidi iz podataka u zadatku). Na drugom mestu unosimo broj stepeni slobode koji je dat u zadatku (za t_1 i t_2 je 18, a za t_3 i t_4 broj stepeni slobode je 998).



Nakon što izračunamo sve kvantile videćemo da su rešenja ovakva:

VAR00001	z1	z2	t1	t2	t3	t4
asdas	2.58	-2.58	2.88	-2.88	2.58	-2.58
	2.58	-2.58	2.88	-2.88	2.58	-2.58
	2.58	-2.58	2.88	-2.88	2.58	-2.58

Uočimo da su vrednosti za z_1 i t_3 odnosno z_2 i t_4 gotovo jednake (odstupaće na nekoj nižoj decimali) što nije slučaj ako uporedimo z vrednosti sa t_1 i t_2 . To se dešava zbog toga što kako povećavamo broj stepeni slobode, odnosno broj ispitanika, to se Studentova distribucija sve više približava standardizovanoj normalnoj i u beskonačnosti se one izjednačavaju (već na 1000 ispitanika t se ponaša kao z). Dakle, sa povećanjem broja stepeni slobode vrednosti kvantila za T -raspodelu približavaju se vrednostima kvantila za standardizovanu normalnu raspodelu. To je razlog zbog kojeg za velike uzorke pri pravljenju intervala poverenja za aritmetičku sredinu umesto kvantila 0.975, odnosno kvantila 0.995 iz Studentove raspodele, možemo koristiti 1.96 (za 95% interval poverenja) i 2.58 (za 99% interval poverenja). Naime, za velike uzorke broj stepeni slobode je veliki pa su kvantili

0.975 i 0.995 za Studentovu raspodelu praktično jednaki 1.96, odnosno 2.58.

Statistička teorema na osnovu koje sve ovo sledi glasi: **Sa povećanjem broja stepeni slobode Studentova raspodela teži standardizovanoj normalnoj raspodeli.**